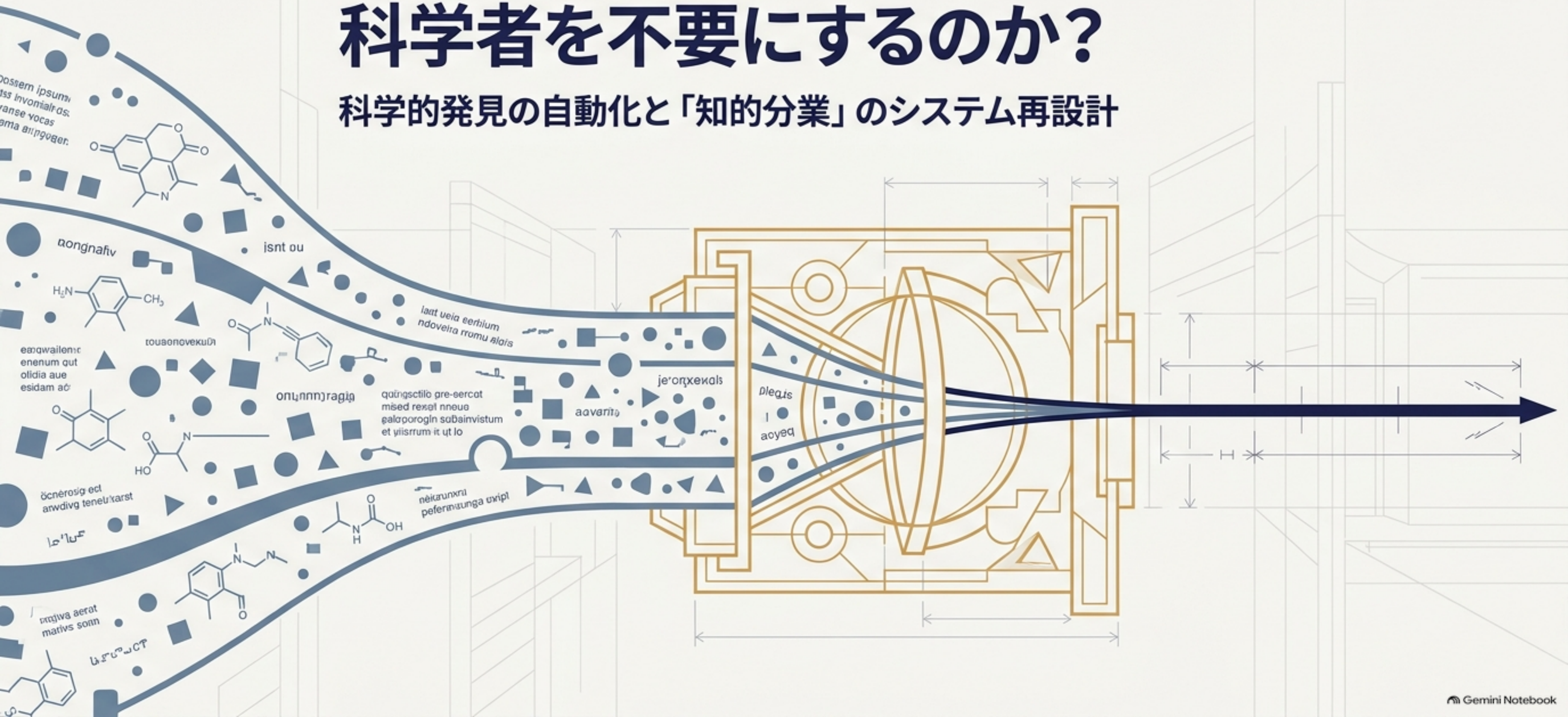


「仮説を作るAI」は 科学者を不要にするのか？

科学的発見の自動化と「知的分業」のシステム再設計



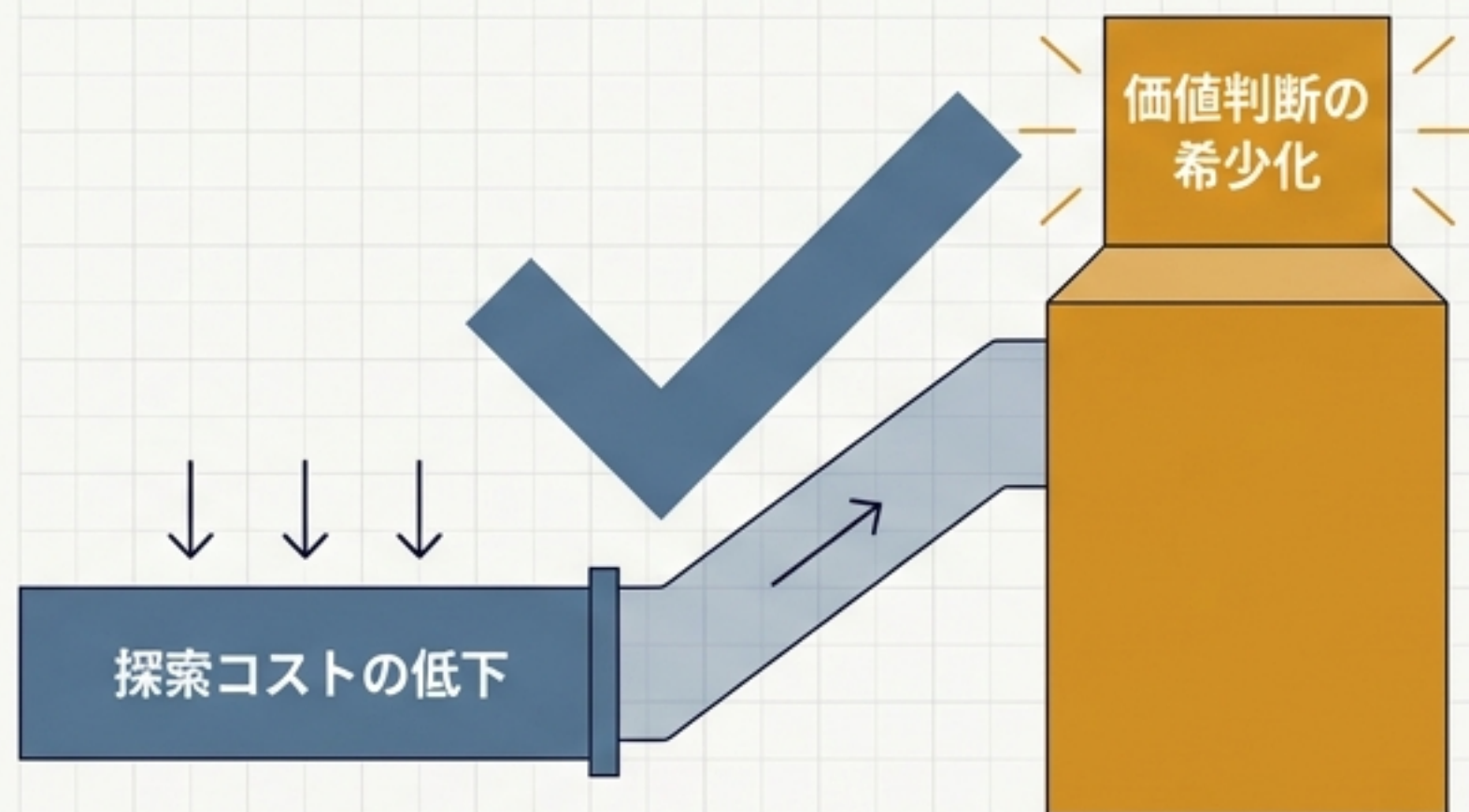
「職業の消滅」ではなく、「ボトルネックの移動」である

Myth (誤った問い)



AIが仮説生成まで行えば、科学者の仕事は奪われる。

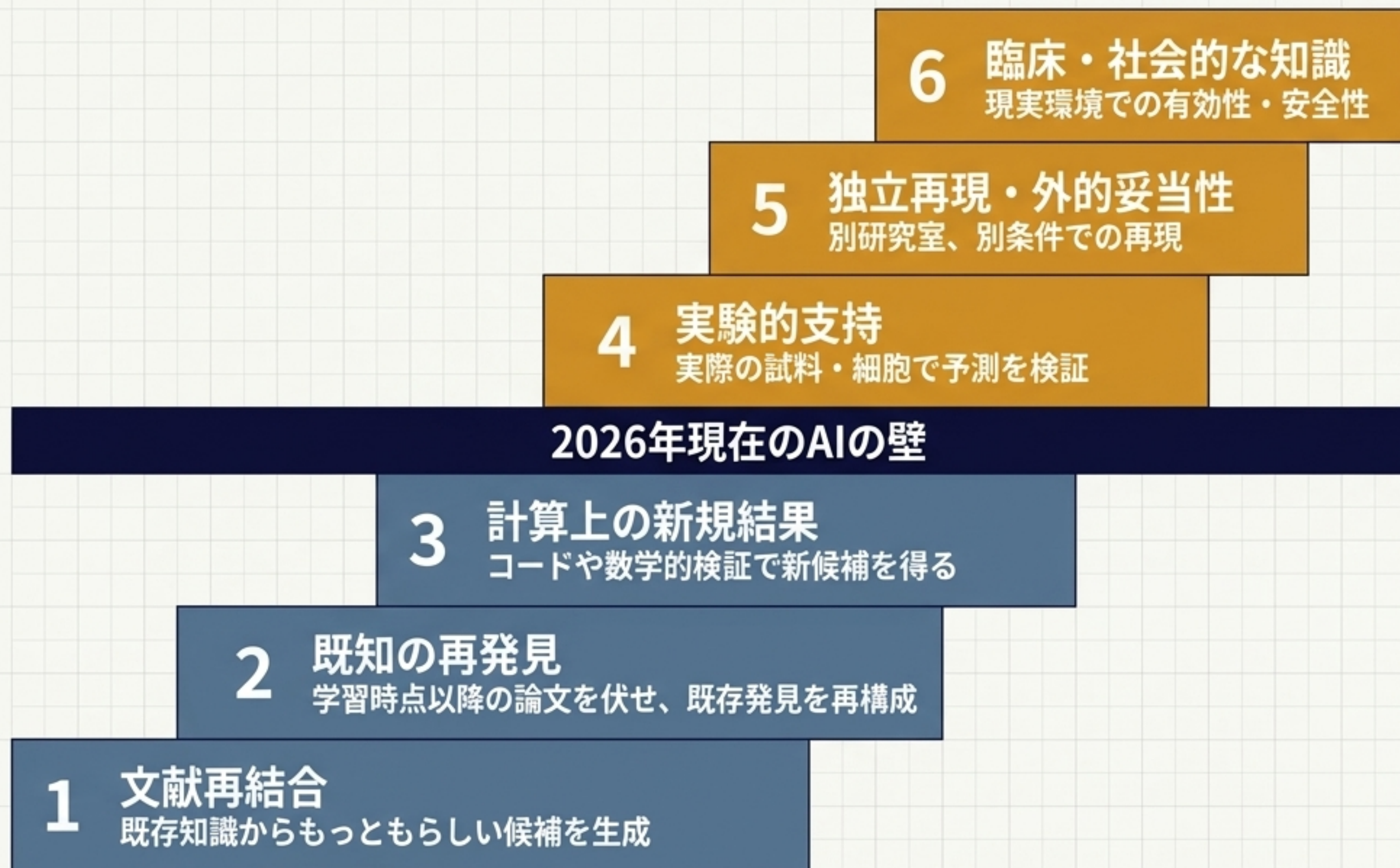
Reality (真の事象)



AIは探索コストを劇的に下げる。結果として、問題設定、測定設計、独立再現といった「人間による検証と価値判断」の希少価値が跳ね上がる。

科学者の中心職能は「アイデアを思いつく人」から「信頼できる知識生成システムを設計・統治する人」へと拡張される。

現在の「AI科学者」は証拠階段のどこにいるのか？

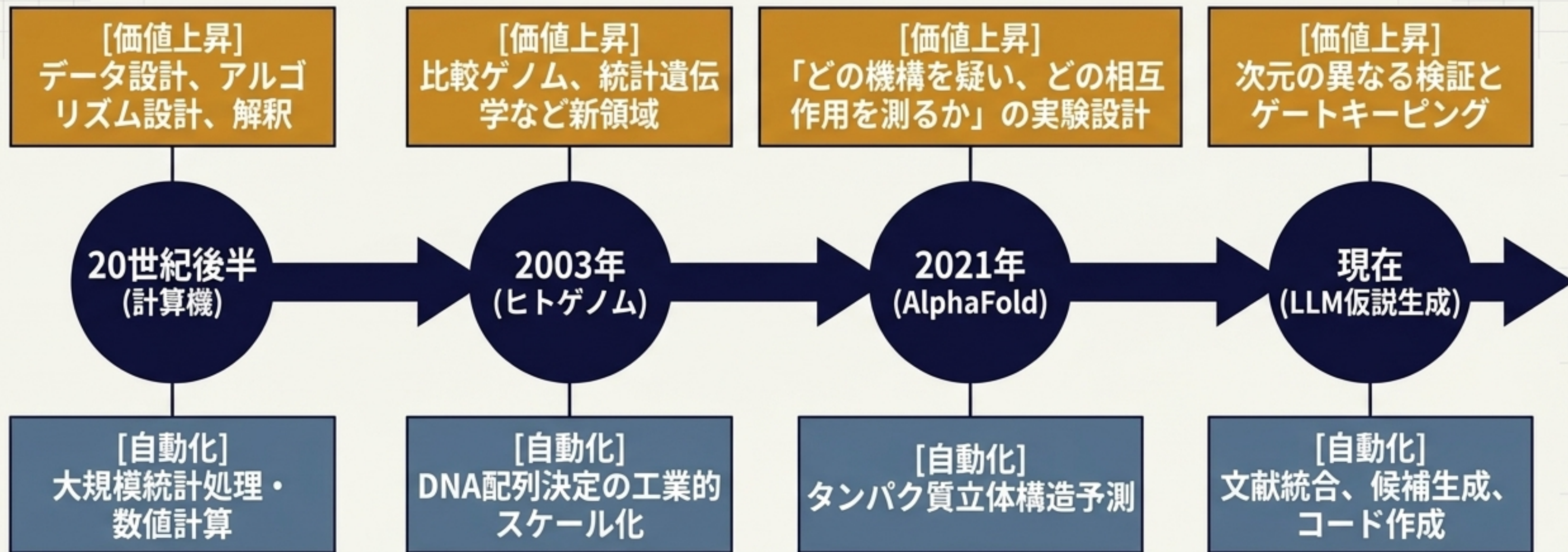


劇的に自動化されているのは「探索空間を広く速く走査すること」であり、「自然界の真理を確立すること」ではない。

最先端「AI科学者」の能力と冷徹な限界

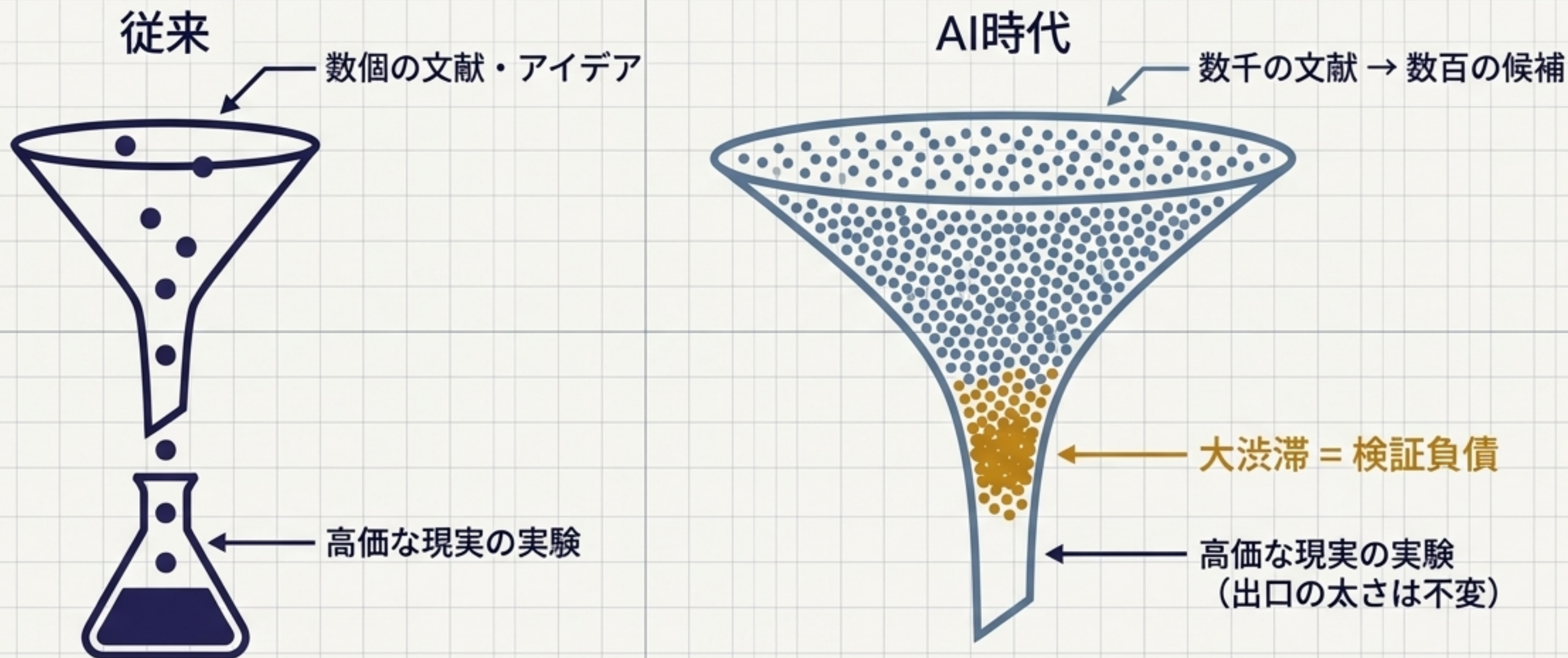
	Google Co-Scientist	The AI Scientist-v2	MOOSE-Chem	HypoBench
自動化の範囲	文献・仮説・実験案	仮説から論文執筆 まで一貫	化学仮説の再発見	仮説生成能力の ベンチマーク
注目すべき実績	AML等でin vitro 検証の一部成功	ICLRワークショップ 平均採択水準超 過	51本の高水準化学 論文をカットオフ 後に再発見	実世界7タスク+合 成5タスクの評価
限界・リスク	幻覚、負の実験結 果の欠落。in vivo 成功は未保証。	開発者自身の本会 議水準評価では3本 中3本が基準未達。	「未知」の発見で はなく、既存論文 への類似度評価。	タスク難化時、人 工データ正解仮説 の回収率38.8%.

歴史的俯瞰：限界費用が下がれば、 次の工程がボトルネックになる



技術がある作業の限界費用を下げると、
その出力を入力として使う「次の工程」の価値が暴騰する。

AI科学時代の最大のリスク：「検証負債（Validation Debt）」



問題はアイデアの不足ではなく、AIが大量生産する偽陽性と検証負債をいかに捌くかという「ゲートキーピング」に移る。

AIと人間の「知的分業」を設計する5つのアーキテクチャ

AIの裁量・探索範囲（小→大）

人間の統治・説明責任のレベル
(直接的統治 → 間接的ゲート制御)

コパイロット型

文献整理や草案作成。
高リスク臨床研究向け。

仮説ポートフォリオ型

数百の候補を生成・ランキン
グ。人間が最終選択。

提案者対監査者型

生成AIと批判AIを戦わせ、
人間が対立する証拠を裁定。

社会的ミッション型

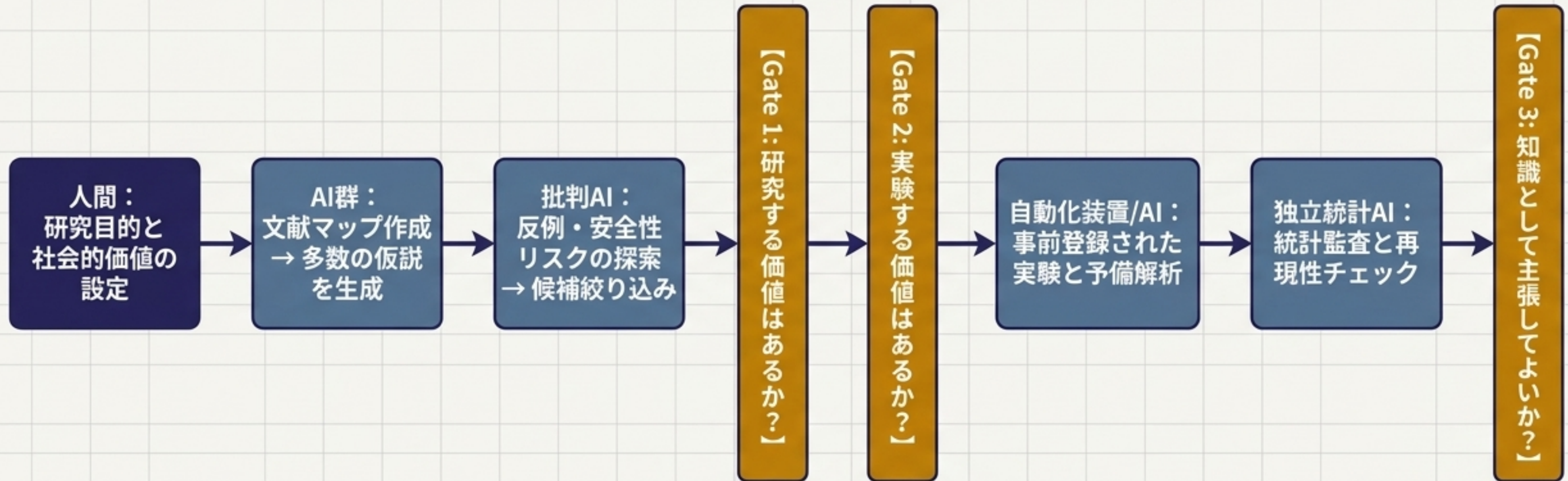
患者や市民が公共目標を与
え、その枠内でAIが探索。

ゲート付き閉ループ型

仮説からロボット実験まで
自動化し、段階移行の
ゲートのみ人間が承認。

境界は固定されない。どの決定に「どの水準の証拠と説明責任が必要か」で分業を設計する。

推奨ワークフロー：「監査者AI」と「人間ゲート」モデル



重要なのはAIの全出力を読むことではない。
ランキングAIの自己増幅を疑い、独立したゲートを管理することである。

AIの自律性を決定する「リスク方程式」

$$\text{許容できるAI自律性} \propto \frac{(\text{結果の機械検証可能性} \times \text{可逆性})}{(\text{失敗時の被害} \times \text{検証困難性})}$$

自律性 高

数学・コード実行

- 答え合わせが一瞬で可能。
失敗してもやり直せる。

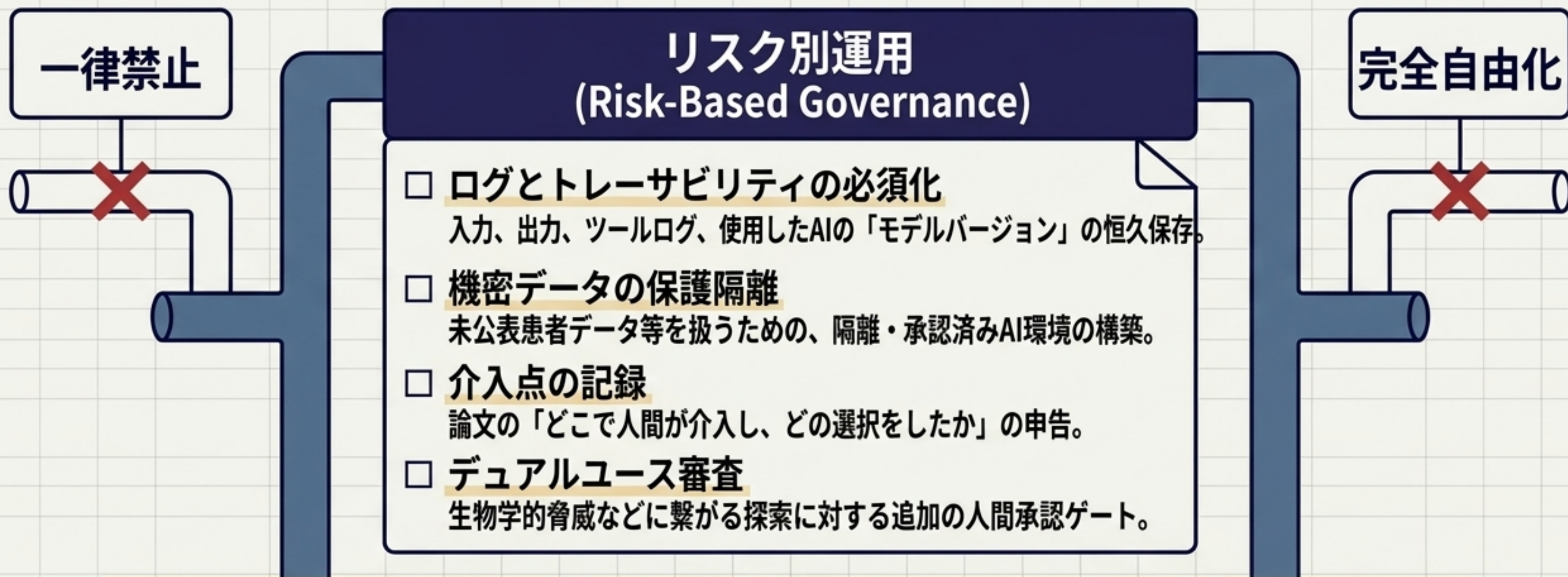
自律性 低・ゲート厚

患者治療・病原体操作

- 致命的ダメージのリスク。
証明に多大なコストがかかる。

「AIを信頼するか」という抽象的論争を捨て、リスクと検証困難性に基づいた実務運用へ。

制度設計のアップデート①：研究室・大学の実務ルール



最終論文だけからAI科学者の「指標誤用」や「データリーク」を見抜くのは不可能に近い。全実行ログの監査体制が必須となる。

制度設計のアップデート ②：研究費審査と学術出版

[過去の評価軸]

派手な新規性、大量の仮
説生成、仮説の数

[未来の評価軸]

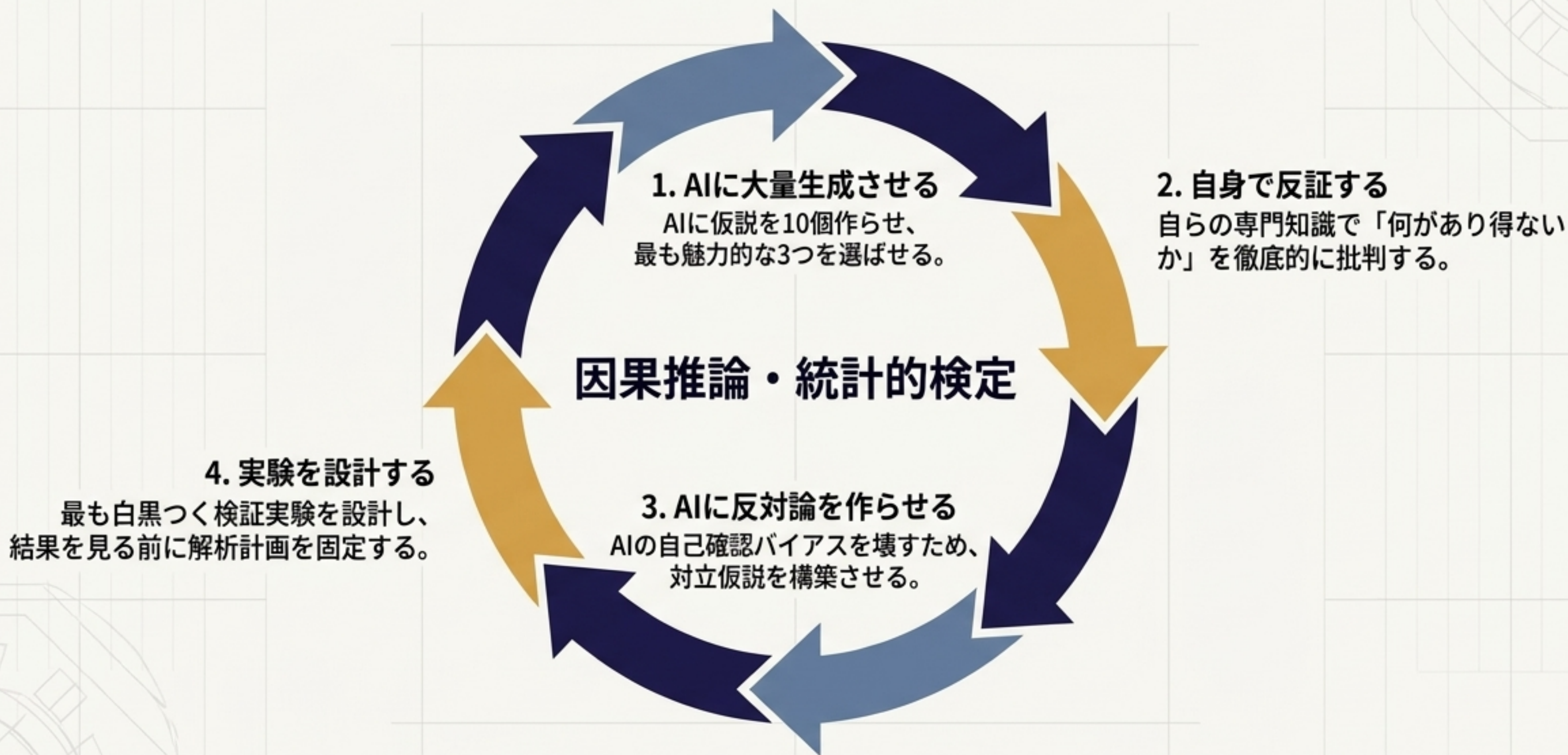
情報価値の高い検証、独立
再現、Negative Results（負
の結果）の共有、高品質な
データセット構築



Nature Portfolio 2026 Standards

- 学術的判断と説明責任はAIへ移転不可。
- AIを用いた査読判断の丸投げは禁止。
- AI Methods Statementの導入: どのようなプロンプト、文献、生成数、選択基準を用いたかの完全開示。

科学教育の再定義：専門知識の役割は「記憶」から「検証」へ



誤った文献まで取り込むAIの幻覚を見抜くためには、かつてなく強い「科学哲学」と「データ生成過程の理解」が必要になる。

結論：AI時代の「科学者」の新しい定義

~~[Old] 優秀な科学者 = 思いついた仮説の質~~

**[New] 優秀な科学者 = 重要な問題の選択 × 探索の多様性 ×
反証力 × 測定品質 × 社会的正統性**

「答え」が安くなった次に、「仮説」も安くなる。これから最も高価になるのは、単なる問いの生成ではない。
「何を知る必要があるかを定め、多数の主体を使って探索し、最も強い反証にさらし、自然界で検証し、
“何を信じるに値する知識として残すか”を決定し、その社会への責任を引き受ける能力」である。
科学者はアイデアの生産者から、知識生成プロセスの設計者・責任者へと進化する。